

Introduction

In today's overly commercialized era, it is crucial for brands to stand out from the information overload and saturated advertising, and to maintain trust with consumers. Meanwhile, more and more brands are incorporating artificial intelligence (AI) into content production and marketing communications to improve efficiency and reduce labor costs (PricewaterhouseCoopers, 2025). However, while virtual influencers offer significant advantages in content production efficiency and brand communication consistency, and their content is often highly scripted and strictly controlled by the brand, their AI-generated identities may still raise questions about authenticity and credibility among consumers.

Patel and Dada (2025) showed that while AI-generated ads and virtual influencers are visually innovative and attractive, some participants felt that the virtual characters were too perfect, creating a sense of unreality, and that virtual influencers failed to evoke emotional resonance and identification. Therefore, even if brands use virtual influencers to briefly grab viewers' attention, maintaining a continuous and stable emotional connection with consumers is paramount from a long-term strategic perspective. Especially in this process, the perceived authenticity of the brand and the level of trust consumers place in it are the foundation for a good relationship between brand and consumers. In this study, perceived brand authenticity is the dependent variable, and AI disclosure is the independent variable, referring to whether brands explicitly inform consumers that content is generated by artificial intelligence. This study will examine the impact of AI disclosure on consumers' perceptions of brand authenticity.

Although existing research has begun to explore the impact of artificial intelligence (AI) information disclosure on brand authenticity, their studies differ significantly, and previous results are inconclusive. Ujjainwala (2025) 's research shows that compared to human influencers, AI influencers significantly reduce consumers' perception of brand authenticity and brand trust, and explicit information disclosure exacerbates this effect. This contrasts sharply with Kowarschik (2025)'s research, which argues that whether or not advertising discloses AI generation does not significantly change consumers' overall perception of brand authenticity (BA) or significantly weaken brand trust.

Building on this, this paper aims to explore whether the moderating effect of AI aversion on this primary influence becomes stronger in advertisements featuring virtual influencers, supplementing existing research conducted by Kowarschik (2025). AI aversion refers to people's existing distrust and aversion to AI algorithms. When people learn that AI is involved in production, they will have more doubts about the AI products than about products produced by humans. Including AI aversion as a moderating variable in AI information disclosure research can account for individual differences. The impact of AI disclosure on brand authenticity, both in strength and direction, may vary depending on consumers' existing level of aversion to AI. This may help explain why prior research had different results, suggesting that the negative effect may vary in strength depending on the contexts and individuals. While existing study has begun to explore the moderating role of AI aversion, Ujjainwala (2025) has limited sample sizes to women. This study employs a more inclusive sample group, observing the responses of consumers of different genders to improve generalizability.

What's more, when encountering highly anthropomorphic virtual influencers, people's AI aversion to AI may influence how individuals perceive the authenticity of the brand, because of the skepticism towards the authenticity of social interactions. As virtual influencers are designed to mimic human communication and build emotional connections, the disclosure may be particularly relevant for how consumers evaluate authenticity. Accordingly, this study aims to answer the following research question:

To what extent does AI disclosure (disclosure vs. non-disclosure) in advertisements created by virtual influencers influence perceived brand authenticity, and to what extent is this effect moderated by AI aversion?

Theoretical Framework

This study aims to expand the theoretical understanding of the relationship between brands and consumers in the field of persuasive communication. Because of inconsistencies in existing research, this study attempts to clarify and explain current academic controversies by introducing AI aversion as a moderation.

With the increasing use of AI-generated content and virtual influencers in advertising, consumers are more often exposed to artificially created content. Some brands explicitly disclose the use of AI, known as AI disclosure. When presenting these advertisements to viewers and consumers, the involvement of AI in the creation process is disclosed to the audience in various ways. For example, this study will use hashtags to convey information about AI's involvement in the creation of the advertisement. In this study, researchers manipulated the independent variable AI disclosure by controlling whether the advertisement contained the label #AI-generated. (□□ method)

While AI disclosure increases information transparency, it may also influence how consumers interpret the authenticity of the advertised content and thus brand authenticity. In Napoli et al.'s (2014) study, brand authenticity is defined as a subjective evaluation of genuineness ascribed to a brand by consumers. Brands can be measured as authentic when they maintain most importantly, are sincere and principled.

However, consumers do not respond to AI-generated content in the same way. This may be due to different AI aversion, which refers to people's negative attitudes towards AI, reflecting a tendency to distrust AI-generated content. Livingston (2025) argues that AI aversion should encompass not only algorithm aversion but also subjective interpretations of the role of artificial intelligence in decision-making. Algorithm aversion specifically refers to the tendency for people to reduce their willingness to use algorithms and instead trust humans, even though algorithms may make the same mistakes as humans (Dietvorst et al., 2015).

Parasocial Interaction Theory

While parasocial relationships initially referred to the relationship between a person and real-life media, recent research by Stein et al. (2022) indicates that in the social media environment, virtual influencers can elicit parasocial interactions (PSI) from users, much like real influencers. Perceived similarity and human-likeness are key factors determining the strength and trust level of PSI.

However, virtual influencers, constructed through algorithms, lack genuine emotional capacity and life experience. Therefore, compared to human influencers, the parasocial relationships they generate are often more limited and fragile. Fandi Omeish et al. (2025) pointed out that AI influencers, constructed through algorithms and lacking genuine emotional capacity and real-life experiences, are considered more difficult to establish parasocial interactions. Especially when consumers become aware of the non-human nature of virtual influencers, this perception may weaken their perceived authenticity and emotional realism, thus undermining parasocial interactions. Empirical research shows that when the artificial origins of virtual influencers are revealed, their perceived humanity decreases, further impacting their parasocial interactions and consequently affecting their credibility and perceived authenticity (Lim & Lee, 2023).

Therefore, while virtual influencers can establish parasocial relationships with consumers, these relationships may be more susceptible to the influence of AI disclosure. When their non-human identities are exposed, the parasocial interaction between consumers and virtual influencers is negatively affected, leading to decreased consumer trust in the virtual influencer and consequently reducing their perception of brand authenticity.

Persuasion Knowledge Model

The Persuasion Knowledge Model posits that after acquiring persuasive knowledge over a long period, people apply it to interpreting, evaluating, and responding to the influence of advertisers and salespeople (Friestad & Wright, 1994). Its characteristic is that consumers use their persuasive knowledge to attempt to analyze whether advertising is manipulative and deceptive, and to understand the persuasive intent behind it, thereby assessing the reliability of the advertising or marketing information they receive. When brands reveal that they use AI to create ads, consumers' knowledge of persuasiveness may be activated, potentially leading them to suspect that the brand's motivation is to generate a perfect image and manipulate virtual influencers to promote a brand's products, rather than to authentically showcase the product. This skepticism triggered by the activation of persuasive knowledge may be more

pronounced in the context of virtual influencers, because virtual influencers have greater controllability and lower subjectivity.

Willemsen et al. (2024) shows that virtual influencers have lower intentions and emotional inclinations compared to regular influencers. When consumers question not only the authenticity of virtual influencer endorsements but also the authenticity of their humanity, they are more likely to doubt their messages because virtual influencers lack genuine emotional experiences and authentic brand preferences. What's more, Jhawar et al. (2023) pointed out that virtual influencers are AI-based, entirely controlled by their parent agency and brands, and have no offline presence. Brands have greater control over virtual influencers, causing them to only convey the messages the brand wants them to, which negatively impacts authenticity and transparency. The high degree of brand control makes it easier for consumers to attribute the content and promotion of virtual influencers to strategic persuasion by brands, rather than recommendations based on genuine experience. In other words, virtual influencers lack independent thought, are unable to oversee brands, and lack subjective initiative.

Research by Baek et al. (2024) indicates that advertising disclosure aims to inform consumers about the persuasive strategies used in advertising; disclosing artificial intelligence as the source of advertising creation may both trigger consumers' persuasive knowledge and simultaneously generate inherent distrust. AI disclosure not only triggers consumers' persuasive knowledge but also amplifies their skepticism because virtual influencers are highly controlled by brands. AI disclosure not only triggers consumers' persuasive knowledge but also amplifies their skepticism because virtual influencers are highly controlled by brands. This is highly consistent with the theoretical expectations of this study, namely that the empirical results of Baek et al. (2024) confirm that AI disclosures can activate persuasive knowledge, prompt consumers to doubt the brand's marketing motives, thereby reducing the authenticity of the advertisement and inducing consumer resistance.

Source Credibility Theory

Meanwhile, Source Credibility Theory argues that the persuasiveness of information depends on the reliability and professionalism of its source (Ohanian, 1990). When ads are labeled as AI-generated, consumers may question the authenticity, reliability, and professionalism of the information conveyed by virtual influencers, believing that virtual influencers cannot filter, process, understand information, or express emotions like human influencers. Therefore, when consumers are explicitly told that an ad is generated by AI, they are more likely to question its authenticity. According to research (Morosoli et al., 2025), disclosing AI-generated content has not increased trust through transparency; instead, it has acted as a negative heuristic, triggering AI aversion among consumers. This is because consumers generally perceive AI as lacking transparency, accountability (because AI operates as a black box), and genuine human emotions. Thus, this hypothesis can be proposed:

H1: Advertisements created by virtual influencers with an AI-generated label lead to lower perceived brand authenticity compared to advertisements without label.

AI aversion as Moderator

Building on this, individual responses to AI disclosures are not uniform, and vary from person to person. These individual differences further influence their judgments about brand authenticity. AI aversion can amplify consumer skepticism and the negative effects of disclosure activation. Algorithm aversion (Dietvorst et al., 2015) suggests that trust in algorithms is more easily eroded, even when AI performs superiorly, people quickly develop distrust after seeing it err. More importantly, Qin et al. (2024) showed that this distrust of algorithms is not uniform across all individuals. The "Capability-Personalization Framework" posits that people's choice to believe in or dislike AI depends on the extent to which AI satisfies their utilitarian and psychological needs in a specific decision-making context. When both needs are met, it depends on the individual's perception of AI's ability to perform the task better within that context. The "Capability-Personalization Framework" proposes two contextual dimensions: perceived capability of AI and perceived necessity for personalization. Each individual's perceived necessity for personalization differs within a given context, leading to variations in individual aversion to AI.

Besides, according to Qiu et al. (2025), individuals with high AI aversion already have low trust in algorithms, making them more susceptible to negative emotions and attitudes

7

towards AI when advertisements reveal they are AI-generated, leading to stronger suspicion about the manipulation and authenticity of the information conveyed by the virtual influencer in the advertisement. Conversely, individuals with lower AI aversion are more accepting of AI, and therefore their evaluation of brand authenticity is relatively less negatively affected when faced with AI disclosures. Empirical research (Wortel et al., 2024) shows that in AI disclosure scenarios, AI aversion amplifies negative individual reactions to brand-related evaluations. Therefore, AI aversion is considered a key moderating variable in this study, it strengthens the negative effect of AI disclosure on perceived brand authenticity. Thus, this hypothesis can be proposed:

H2: Advertisements created by virtual influencers with an AI-generated label lead to lower perceived brand authenticity compared to advertisements without label, and this effect is stronger for individuals with high AI aversion compared to individuals with low AI aversion.

Conceptual Model

Based on the theoretical framework and hypothesis established above, this conceptual model provides a visual presentation of the key relationships examined in this research.

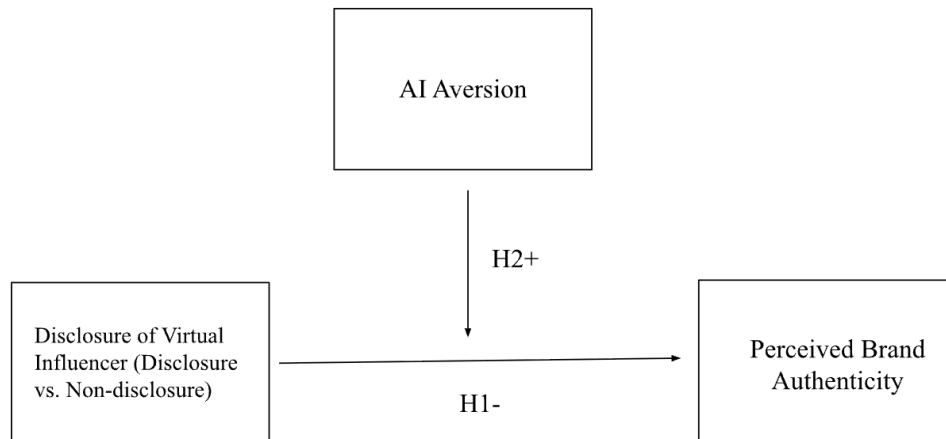


Figure1. *Conceptual model*

8

Research Design

Method

This study uses a mixed experimental design. The experiment was designed to integrate the requirements of four related research projects. However, this study only focuses on the between-subjects component of the whole design, which consisted of a one-factor experimental design. The independent variable, AI disclosure will be manipulated by presenting participants with stimulus material either with or without AI-generated label. And the dependent variable, participants' perceived brand authenticity will be measured to examine the impact of the disclosure. In addition, AI aversion is introduced as a moderator. During the experiment, demographic variables (gender, age and education) are measured and included as control variables. This allows control of confounding effects and ensures that differences in perceived authenticity are determined solely by the disclosure of the AI, which enhances internal validity.

In this mixed-design laboratory experiment (N = 128), we manipulated the AI disclosure the participants exposed (i.e., AI-generated labeled vs. non labeled) as a between-subjects factor and included the context of influencer (i.e., sports advice influencer vs. sexual services influencer) as a within-subjects factor.

An experimental approach was considered most suitable for this research, because only experiment allows for causal inference by systematically manipulating AI disclosure while controlling other factors constant (Shadish et al., 2002). Random assignment is used to assign participants to the different conditions, which enhances internal validity and balances confounding elements across groups (reference).

Because the order of the two treatments in this experiment is randomized, to some extent, the interference between treatments is reduced, thus reducing the order effects and carryover effects caused by fatigue and sequence (Hedayat & Kirk, 1970). Hedayat and Kirk (1970) discusses demand characteristics as unintentional cues given to participants in the experiment that hint at the researchers' desired outcomes. Building on this, between-subjects designs may reduce demand characteristics because participants are only exposed to only one condition

9

(i.e., AI-generated labeled vs. non labeled) and therefore have fewer cues to infer are unaware of other conditions, which means lower risk of participants bias and awareness.

Sample

128 participants were recruited from the WhatsApp and WeChat group chat through the online survey platform Qualtrics. A convenience sampling method will be used to collect data. Following the suggestion of Faul et al. (2009), we estimated the required sample size using G*Power. Based on the between-group design, with $\alpha = 0.05$, power = 0.80, and a moderate effect size ($f^2 = 0.15$), this study requires 68 participants. To ensure sufficient power, this study aimed to recruit at least 128 participants.

The data of 4 participants were excluded from the analyses because they are not European residents or they are under 18 years old. The final sample size of this study was 124 participants ($M_{age} = 0$, $SD_{age} = 0$, 0 male, 0 female, 0 prefer not to say, 0 non-binary / third gender). Participants were randomly assigned to two of two conditions, they were either shown posts from labeled sports advice influencers and then labeled sexual services influencers ($n =$), or posts from

unlabeled sports advice influencers and then unlabeled sexual services influencers (n=), or posts from labeled sexual services influencers and then labeled sports advice influencers (n=), or posts from unlabeled sexual services influencers and then unlabeled sports advice influencers (n=).

Procedure

Participants were first invited through an online link for the study, which was placed on WhatsApp and WeChat group chat and received a €2 bonus. After that, the researchers randomly selected participants in the new library of the University of Amsterdam to fill out a questionnaire via an online link and receive a chocolate after completing it.

All participants were provided with informed consent before the questionnaire began. Next, participants were shown a set of questions measuring the moderator, their existing aversion to AI. On the same page, they were shown another set of questions about their attitude towards online influencers as fillers to eliminate participants' assumptions about the

10

purpose of this questionnaire and assumed in advance that the people in Instagram posts without AI-generated labels are virtual influencers. Later on, participants answered demographic questions, including information such as age, gender, and education level.

Participants were randomly assigned to four groups. One group was shown labeled Instagram posts from a sports advice influencer first, followed by labeled Instagram posts from a sexual services influencer. The second group was shown unlabeled Instagram posts from a sports advice influencer first, followed by unlabeled Instagram posts from a sexual services influencer. The third group was shown labeled Instagram posts from a sexual services influencer first, followed by labeled Instagram posts from a sports advice influencer. The fourth group was shown unlabeled Instagram posts from a sexual services influencer first, followed by unlabeled Instagram posts from a sports advice influencer.

Participants who were randomly assigned to the first group viewed an Instagram story post by a sports advice influencer and were required to watch it for at least 20 seconds to ensure they did not skip the stimulus material. After viewing the stimulus material, they first completed a manipulation check, asking if the post contained AI-generated labels. They then completed a series of questions that measured their content perception, immediate reaction to the post, perceived brand authenticity, perceived influencer authenticity, anthropomorphism, and their purchase intention. And then they were shown to another Instagram story post by a sexual services influencer. After that, they completed the manipulation check again and the attention check to ensure that participants maintained their attention. In addition, there are questions similar to those mentioned earlier regarding measuring their content perception and immediate reaction to the post in the next.

The other three conditions follow the same logic. At the end, participants were debriefed and thanked for the participation.

Stimulus Materials

All participants saw two Instagram story posts, each with two different versions. Both versions use the same image. One version uses a colored label on the upper half of the image that reads *"This human figure is AI generated"*, while the other version does not have this

11

label, thus enabling the manipulation of AI disclosure. And all other elements remained the same.

Participants will only be shown two labeled Instagram story posts or two unlabeled Instagram story posts. Of the two Instagram story posts that participants will see, one is a selfie taken by Alex at the gym, accompanied by the fitness course brand promotion information *"my personal fitness training course is here and here only"* and a course link *@AlexpMuscle*. The other post is a selfie taken by Fox at the gym from a similar angle, accompanied by his personal account information *"my content is here and here only"* and his personal account website link *@OnlyFans*. Aside from differences in people and account content, two Instagram story posts were matched in visual composition, layout, title structure,

accessories and clothing to ensure that labels were the only difference between the conditions.

Measures

AI Aversion. AI aversion refers to people's negative attitudes towards AI, reflecting a tendency to distrust AI-generated content (). We included a 5-point likert scale used in previous research that measured AI Aversion ranging from 1 (*strongly disagree*) to 5 (*strongly agree*). This scale includes four items to measure general attitudes toward artificial intelligence, namely "I believe that artificial intelligence will improve my life", "I believe that artificial intelligence will improve my work", "I think I will use AI technologies in the future", and "I think AI technology is beneficial for humanity" (Grassini, 2023). Because these questions all reflect a positive AI attitude, reverse coding was used when processing the data, so that higher scores reflected greater AI aversion.

An exploratory factor analysis was conducted to examine whether the four items formed a single factor. The results indicated that all items loaded on one factor (eigenvalue = ?) explaining ?% of the variance. Sampling adequacy was supported by a ? high KMO value (.), and Bartlett's test of sphericity was ? significant ($\chi^2(10) = ?$, $p = ?$), confirming the suitability of the data for factor analysis. The results of reliability analysis showed that the items formed a ? reliable scale (Cronbach's Alpha = ?). Participants' responses to the four items were averaged to create an overall score. The mean score was $M = ?$ ($SD = ?$).

12

Perceived Brand Authenticity. Perceived brand authenticity refers to a subjective evaluation of genuineness ascribed to a brand by consumers (Napoli et al., 2014). We included a 5-point likert scale used in previous research that measured Perceived Brand Authenticity ranging from 1 (*strongly disagree*) to 5 (*strongly agree*). This scale includes five items to measure sincerity dimension from Napoli et al. (2014)'s Brand Authenticity Framework, namely "This brand is sincere", "This brand is honest in its communication", "This brand does not pretend to be something it is

not”, “This brand is genuine”, and “This brand stays true to its values”.

This study did not use the other two dimensions in Napoli et al.'s (2014) Brand Authenticity Framework, quality commitment and heritage. This is because AI labels essentially influence people's perceptions of a brand's authenticity, and this study does not focus on people's perceptions of brand quality. Furthermore, the brand selected for this study was a personal fitness course brand with no brand history, and participants could not determine whether its heritage had been maintained.

An exploratory factor analysis was conducted to examine whether the four items formed a single factor. The results indicated that all items loaded on one factor (eigenvalue = ?) explaining ?% of the variance. Sampling adequacy was supported by a ? high KMO value (.), and Bartlett's test of sphericity was ? significant ($\chi^2(10) = ?$, $p = ?$), confirming the suitability of the data for factor analysis. The results of reliability analysis showed that the items formed a ? reliable scale (Cronbach's Alpha = ?). Participants' responses to the five items were averaged to create an overall score. The mean score was $M = ?$ ($SD = ?$).

Results

First, data from participants who failed the attention check will be removed. Specifically, we will use frequency tables to check the missing value and examine the control variables: gender, age, education, and familiarity with the virtual influencer. The mean, standard deviation, and distribution of brand authenticity and AI aversion will be inspected. A chi-square test will be used for manipulation checks. Randomisation checks will be

13

conducted to examine whether control variables differ between two conditions. Chi-square will be used to examine categorical control variables, namely gender and education.

Description and Manipulation check

The descriptive statistics (frequency, mean value and standard deviation) will be used to examine the distribution of all variables. Means and standard deviations were computed for all core constructs (Appendix ?/ Table ?). Perceived Brand Authenticity (PBA) (M=?; SD=?) and AI Disclosure (AD) (M=?; SD=?) were rated highest/lowest, indicating generally positive/negative/neutral responses. AI Aversion(AA) (M=?; SD=?) was near the top/bottom/midpoint, suggesting positive/negative/neutral perceptions.

Reliability analyses (Cronbach's alpha) were conducted for the brand authenticity scale and the AI aversion scale.

A chi-square test was used to examine whether the manipulations were successful for this experiment. Results are presented in Table?. The result showed that participants exposed to the labeled/unlabeled condition reported higher authenticity (M = ?, SD = ?) than those who exposed to labeled/unlabeled condition (M = ?, SD = ?), $t(?) = ?$, $p ?$, 95% CI [?, ?], $d = ?$, indicating that manipulation was successful/unsuccessful and strong/weak/moderate.

Testing hypothesis

Finally, to examine whether AI disclosure negatively influences PBA, Model 1 of PROCESS macro version conducted to test the main effect of AI disclosure to perceived brand authenticity. The interaction effect between AI disclosure and AI aversion will be examined to test moderation effect. Results of PROCESS are presented in Tabel?.

Main effect

Results of PROCESS macro version showed that AI disclosure has a significant/non-significant effect on perceived brand authenticity, indicating that individuals in labeled condition perceived more or less authenticity than those in unlabeled conditions/there is no relationship between AI disclosure and perceived brand authenticity.

Moderation effect

Subsequently, results of PROCESS macro version showed that the interaction effect between AI disclosure and AI aversion was significant/non-significant, $b = ?$, $SE = ?$, $p = ?$, 95% CI $[?, ?]$, indicating that individuals with higher AI aversion perceived more/less authenticity in labeled conditions/AI aversion did not significantly moderate the relationship between AI disclosure and perceived brand authenticity.

Conclusion and Discussion

To summarize, this study examined the effect of AI disclosure (disclosure vs. non-disclosure) in advertisements created by virtual influencers on perceived brand authenticity, and further investigated the moderating role of AI aversion in affecting this relationship. Overall, the findings revealed a significant/not significant interaction between the disclosure of the virtual influencer's origin and perceived brand authenticity.

Limitation

The findings of this study contribute to understanding how AI disclosures in virtual influencer advertising affect the perception of brand authenticity. However, some limitations exist, which can provide some reference suggestions for future research. To begin with, this study used convenience sampling, which may limit generalizability of the findings. Participants were recruited through online social media, which potentially results the sample is not representative of the broader population. Moreover, The sample mainly consists of young, highly educated individuals who are familiar with social media. This group may have a higher acceptance of virtual influencers and artificial intelligence, which limits the applicability of the research results to other populations. Future studies should include more diverse samples with different cultural backgrounds and different age groups. This will help examine whether the impact of AI-driven information disclosure on consumers' perceived brand authenticity is consistent across a broader consumer group, or whether this impact differs between younger consumers who are more receptive to new technologies and older consumers. Additionally, the stimulus material used in this

study was selected from photos of real human influencers, even those labeled as AI-generated. However, most virtual influencers in real life are not entirely like real people and usually have some unnatural AI elements. This may reduce ecological validity. Further, another limitation of this study is the use of a mixed experimental design, which may lead to order effect and carryover effect.

15

Participants were exposed to two stimulus materials in a random order. If participants are first exposed to stimulus material containing sexual context, and then to stimulus material containing sport context used in this study, participants may feel uncomfortable and averse to the material containing sexual services, which may affect their judgment of the latter material. Plus, after being exposed to the first stimulus material, participants are asked whether there are AI-generated labels in the images to check for manipulation. This may make participants aware of the purpose of the study and consciously answer the questions following the next material, thereby amplifying the effect of AI disclosure in the results. Future research could use between-subjects design to avoid order effect and carryover effect, thereby more accurately detecting causal relationships between variables.

References

Baek, T. H., Kim, J., & Kim, J. H. (2024). Effect of disclosing AI-generated content on prosocial advertising evaluation. *International Journal of Advertising*, 1-22.

<https://doi.org/10.1080/02650487.2024.2401319>

Brüns, J. D., & Meißner, M. (2024). Do You Create Your Content yourself? Using Generative Artificial Intelligence for Social Media Content Creation Diminishes Perceived Brand Authenticity. *Journal of Retailing and Consumer Services*, 79(0969-6989), 103790-103790. <https://doi.org/10.1016/j.jretconser.2024.103790>

Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015a). Algorithm Aversion: People Erroneously Avoid Algorithms after Seeing Them Err. *SSRN Electronic Journal*, 144(1).

<https://doi.org/10.2139/ssrn.2466040>

Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015b). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126.

16

Fandi Omeish, Shaheen, A., Sager Alharthi, & Aisyah Alfaiza. (2025). Between human and AI influencers: parasocial relationships, credibility, and social capital formation in a collectivist market: a study of TikTok users in the Middle East. *Discover Sustainability*, 6(1). <https://doi.org/10.1007/s43621-025-00891-w>

Faul, F., Erdfelder, E., Buchner, A., & Lang, A.-G. (2009). Statistical Power Analyses Using G*Power 3.1: Tests for Correlation and Regression Analyses. *Behavior Research Methods*, 41(4), 1149–1160. <https://doi.org/10.3758/brm.41.4.1149>

Friestad, M., & Wright, P. (1994). The persuasion knowledge model: How people cope with persuasion attempts. *Journal of Consumer Research*, 21(1), 1–31. <https://doi.org/10.1086/209380>

Grassini, S. (2023). Development and validation of the AI attitude scale (AIAS-4): A brief measure of general attitude toward artificial intelligence. *Frontiers in Psychology*, 14(1), 1191628. <https://doi.org/10.3389/fpsyg.2023.1191628>

Hedayat, A., & Kirk, R. E. (1970). Experimental Design: Procedures for the Behavioral Sciences. *Biometrics*, 26(3), 590. <https://doi.org/10.2307/2529119>